

Single Bit and Reduced Dimension Diffusion Strategies Over Distributed Networks

Muhammed O. Sayin, Suleyman S. Kozat*, *Senior Member, IEEE*

Abstract—We introduce novel diffusion based adaptive estimation strategies for distributed networks that have significantly less communication load and achieve comparable performance to the full information exchange configurations. After local estimates of the desired data is produced in each node, a single bit of information (or a reduced dimensional data vector) is generated using certain random projections of the local estimates. This newly generated data is diffused and then used in neighboring nodes to recover the original full information. We provide the complete state-space description and the mean stability analysis of our algorithms.

Index Terms—Diffusion, distributed, single-bit, compressed.

I. INTRODUCTION

DISTRIBUTED adaptive estimation utilizes a network of nodes that observe a monitored phenomena with different view points. This broadened perspective can be used to enhance estimation performance or eliminate obstructions in the environment, which may not be achieved using a single node [1]. The distributed algorithms usually target to reach the best estimate that could be produced when the individual nodes have access to all observations across the whole network. However, there is naturally a trade-off between the amount of cooperation and required communication among the nodes [1], [2].

The diffusion based distributed algorithms define a strategy in which the nodes from a predefined neighborhood could share information with each other [1], [2]. Such approaches are stable against time-varying statistical profiles [1], however, require a high amount of communication resources. For example, in a network of N nodes, where $\bar{n}$ denotes the average number of nodes in a neighborhood, then $N \times \bar{n}$ number of parameter estimates should be communicated among nodes on the average at each time.

In this letter, we propose diffusion based cooperation strategies that have significantly less communication load (e.g., a single bit of information exchange) and achieve comparable performance to the full information exchange configurations under certain settings. In this framework, after local estimates of the desired vector is produced in each node, a single bit of information (or a reduced dimensional data vector) is generated using certain random projections of the local estimates. This new information is diffused and used in neighboring nodes instead of the original estimates; significantly reducing the communication load in the network. We only require synchronization of this randomized projection

operation, which can be achieved using simple pilot signals [3]. Note that our approach differs from quantization based diffusion strategies such as [4] in terms of the compression of the diffused information. In [4], a quantized parameter estimate is exchanged among nodes to avoid infinite precision in the transmission. In [5], the sign of the innovation sequence in the context of decentralized estimation and in [6], the relative difference between the states of the nodes in a consensus network are exchanged using a single bit of information. Here, we substantially compress the exchanged information, even to a single bit, and perform local adaptive operations at each node to recover the full information vector. In this sense, our method is more akin to compressive sensing rather than to a quantization framework.

Our main contributions include: 1) We propose algorithms to significantly reduce the amount of communication between nodes for diffusion based distributed strategies; 2) We analyze the stability of the algorithms in the mean under certain statistical conditions; 3) We illustrate the comparable convergence performance of these algorithms in different numerical examples. We emphasize that although we only provide the mean stability analysis due to space limitations, the mean-square convergence and tracking analysis are carried out in a similar fashion (following [1]) in a separate paper submission.

The letter is organized as follows. In Section II, we introduce the framework and the studied problem. The new approaches are derived in Section III. In Section IV, we analyze the mean stability of our approaches. Numerical examples and concluding remarks are provided in Section V.

II. PROBLEM DESCRIPTION

Consider the widely studied spatially distributed framework [1], [2]. Here, we have N number of nodes where two nodes are considered neighbors if they can exchange information. For a node i , the set of indexes of its neighbors including the index of itself is denoted by $\mathcal{N}_i$. At each node, an unknown desired vector¹, $\mathbf{w}_o \in \mathbb{R}^m$, is observed through a linear model $d_i(t) = \mathbf{w}_o^T \mathbf{u}_i(t) + v_i(t)$,² assuming the observation noise is temporally and spatially white (or independent), i.e., $E[v_i(t)v_j(l)] = \sigma_i^2 \delta(i-j)\delta(t-l)$, where $\delta(\cdot)$ is the Kronecker delta and σ_i^2 is the variance of the noise. The regression vectors

¹Although we assume a time invariant desired vector, our derivations can be readily extended to certain non-stationary models [3].

²We represent vectors (matrices) by bold lower (upper) case letters. For a matrix $\mathbf{A}$ (or a vector $\mathbf{a}$), $\mathbf{A}^T$ is the transpose. $\|\mathbf{a}\|$ is the Euclidean norm. For notational simplicity we work with real data and all random variables have zero mean. The sign of a is denoted by $\text{sign}(a)$ (0 is considered positive without loss of generality). For a vector $\mathbf{a}$, $\dim(\mathbf{a})$ denotes the length. The expectation of a vector or a matrix is denoted with an over-line, i.e. $E[\mathbf{a}] = \bar{\mathbf{a}}$. The $\text{diag}(\mathbf{A})$ returns a new matrix with only the main diagonal of $\mathbf{A}$ while $\text{diag}(\mathbf{a})$ puts $\mathbf{a}$ on the main diagonal of the new matrix.

This work is supported in part by TUBITAK, Contract No:112E161. Muhammed O. Sayin (m_sayin@ug.bilkent.edu.tr) and Suleyman S. Kozat (kozat@ee.bilkent.edu.tr) are with the Department of Electrical and Electronics Engineering, Bilkent University, Ankara, Turkey (Tel: +90312 2902336 and Fax: +90312 2901223)

are also assumed to be spatially and temporally uncorrelated with each other and with the observation noise. At each node an adaptive estimation algorithm is working such as the LMS algorithm [3] given as

$$\phi_i(t+1) = (\mathbf{I} - \mu_i \mathbf{u}_i(t) \mathbf{u}_i^T(t)) \mathbf{w}_i(t) + \mu_i d_i(t) \mathbf{u}_i(t),$$

$\mu_i > 0$. As the diffusion strategy, we use the adapt-then-combine (ATC) diffusion strategy as an example since it is shown to outperform the combine-then-adapt diffusion and consensus strategies under certain conditions [7]. However, our derivations also cover these distributed strategies. In the ATC strategy, at each node i , the final estimate is constructed as

$$\mathbf{w}_i(t+1) = \sum_{k \in \mathcal{N}_i} \lambda_{i,k} \phi_k(t+1),$$

where $\lambda_{i,k}$'s are the combination weights $\sum_{k \in \mathcal{N}_i} \lambda_{i,k} = 1$ and $\lambda_{i,k} \geq 0$. The combination weights $\lambda_{i,k}$ can also be adapted in time, affinely constrained or unconstrained [8]. We stick to constant-in-time weights with the simplex constraint since the stabilization effect of such weights is demonstrated in [1].

In the diffusion based distributed networks, whole parameter estimates are exchanged within the neighborhood. In the next section, we introduce two different approaches in order to reduce the amount of information exchange between nodes.

III. NEW DIFFUSION STRATEGIES

A. Reduced Dimension Diffusion

In the first approach, each node calculates a reduced dimensional vector through a linear transformation $\mathbf{z}_k(t+1) = \mathbf{C}(t+1) \phi_k(t+1)$, where $\dim(\mathbf{z}_k(t+1)) \ll \dim(\phi_k(t+1))$, and transmits $\mathbf{z}_k(t+1)$ instead of $\phi_k(t+1)$. We use a randomized linear transformation matrix $\mathbf{C}(t+1)$ where the size of the matrix determines the compression amount. Each neighboring node uses the same $\mathbf{C}(t+1)$. After receiving $\mathbf{z}_k(t+1)$, a neighbor node i constructs an estimate $\mathbf{a}_k(t+1)$ of the original $\phi_k(t+1)$ using a *minimum disturbance criteria* [3] as

$$\mathbf{a}_k(t+1) = \arg \min_{\mathbf{a}} \|\mathbf{a} - \mathbf{a}_k(t)\| \quad (1)$$

$$\text{such that } \mathbf{C}(t+1) \mathbf{a} = \mathbf{z}_k(t+1),$$

where $\mathbf{a}_k(t) \in \mathbb{R}^m$. Note that (1) yields the NLMS algorithm [3] as

$$\mathbf{a}_k(t+1) = \mathbf{a}_k(t) + \sigma_k \mathbf{C}(t+1)^T [\mathbf{C}(t+1) \times \mathbf{C}(t+1)^T]^{-1} (\mathbf{z}_k(t+1) - \mathbf{C}(t+1) \mathbf{a}_k(t)), \quad (2)$$

where a learning rate $\sigma_k > 0$ is also incorporated after (1). After $\mathbf{a}_k(t+1)$'s are calculated, we construct the final estimate at node i as

$$\mathbf{w}_i(t+1) = \lambda_{i,i} \phi_i(t+1) + \sum_{k \in \mathcal{N}_i \setminus i} \lambda_{i,k} \mathbf{a}_k(t+1). \quad (3)$$

Remark 1: For a time invariant projection matrix, $\mathbf{C}(t) = \mathbf{C}$, the exchanged estimate $\mathbf{a}_k(t)$ converges to the projection of the original parameter estimate $\phi_k(t)$ onto the column space of the matrix $\mathbf{C}$ (provided that adaptation is fast enough). In order to avoid biased convergence, we choose randomized projection matrices that span the whole parameter space.

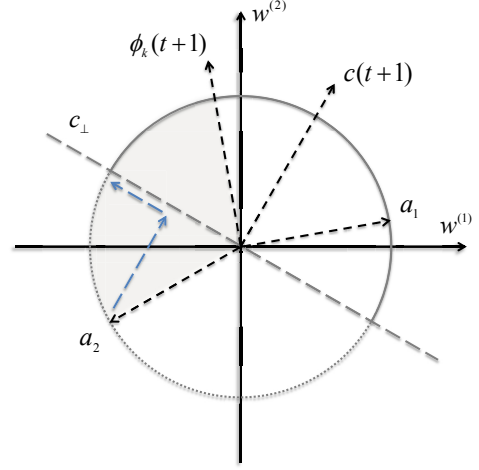


Fig. 1: Single bit diffusion in two dimensions, i.e., $\mathbf{w}_o \in \mathbb{R}^2$. As an example, one can have $\mathbf{a}_k(t) = \mathbf{a}_1$ or $\mathbf{a}_k(t) = \mathbf{a}_2$. $\mathbf{c}_\perp$ denotes the vector space perpendicular to $\mathbf{c}(t+1)$ in two dimensions and the shaded area represents the update region for $\mathbf{a}_k(t) = \mathbf{a}_2$

Remark 2: One can also use an ordinary LMS update to train $\mathbf{a}_k(t)$ to avoid the inversion operation in (2), considering $\mathbf{z}_k(t+1)$ as the desired data and $\mathbf{C}(t+1)$ as the regression matrix. However, since the dimensions of $\mathbf{z}_k(\cdot)$'s are much smaller than the dimension of $\mathbf{w}_o$, e.g., in our simulations we use scalar $z_k(\cdot)$'s with $m = 1$, one can use the NLMS update for $\mathbf{a}_k(\cdot)$'s without significant computational increase.

In the following, we further reduce the amount of transmitted information by diffusing a single bit of information instead of a scalar.

B. Single Bit Diffusion

In this approach, we exchange only the sign of the linear transformation $\mathbf{z}_k(t+1) = \mathbf{c}(t+1)^T \phi_k(t+1)$. According to the transmitted sign, the neighboring node i can construct an estimate $\mathbf{a}_k(t+1)$ of $\phi_k(t+1)$ as

$$\mathbf{a}_k(t+1) = \arg \min_{\mathbf{a}} \|\mathbf{a} - \mathbf{a}_k(t)\| \text{ such that } \quad (4)$$

$$\text{sign}(\mathbf{c}(t+1)^T \mathbf{a}) = \text{sign}(z_k(t+1)) \text{ and } \quad (5)$$

$$\|\mathbf{a}\| = 1. \quad (6)$$

To solve (4), we observe from Fig. 1 that we can only have two different cases for $\mathbf{a}_k(t)$. In the first case, we have $\text{sign}(\mathbf{c}(t+1)^T \mathbf{a}_k(t)) = \text{sign}(z_k(t+1))$, e.g., $\mathbf{a}_k(t) = \mathbf{a}_1$ in the figure. In this case, no update is needed, $\mathbf{a}_k(t+1) = \mathbf{a}_k(t)$, since $\mathbf{a}_k(t)$ satisfies both conditions (5), (6) and $\|\mathbf{a}_k(t+1) - \mathbf{a}_k(t)\| = 0$. In the second case, we have $\text{sign}(\mathbf{c}(t+1)^T \mathbf{a}_k(t)) \neq \text{sign}(z_k(t+1))$, e.g., $\mathbf{a}_k(t) = \mathbf{a}_2$ in the figure. For this case, i.e., $\mathbf{a}_k(t) = \mathbf{a}_2$, we only need to project $\mathbf{a}_k(t)$ to the half hyper sphere (shown as a half circle in two dimensions in Fig. 1), which corresponds to the constraints (5) and (6). This projection can be readily accomplished by first projecting $\mathbf{a}_k(t)$ to the vector space perpendicular to $\mathbf{c}(t+1)$ and then scaling the projected vector to have unit norm. This yields the following

$$\mathbf{a}_k(t+1) = \frac{\mathbf{a}_k(t) - \gamma(t+1) \mathbf{c}(t+1) \mathbf{c}(t+1)^T \mathbf{a}_k(t)}{\|\mathbf{a}_k(t) - \gamma(t+1) \mathbf{c}(t+1) \mathbf{c}(t+1)^T \mathbf{a}_k(t)\|}$$

update, where

$$\gamma(t+1) \triangleq \frac{1 - \text{sign}(z_k(t+1)) \text{sign}(\mathbf{c}(t+1)^T \mathbf{a}_k(t))}{2 \mathbf{c}(t+1)^T \mathbf{c}(t+1)}.$$

Here, (6) is needed to resolve the inherent amplitude uncertainty in (5) since the diffused sign bit does not carry any amplitude information. Without (6), the amplitude of the constrained estimates diminishes to zero by updates. We resolve the amplitude uncertainty in the final combination by multiplying the unit norm estimate $\mathbf{a}_k(t+1)$ with the magnitude of the local parameter estimate $\phi_i(t+1)$. This scaling with the norm of $\phi_i(t+1)$ results in the rotated parameter estimation in the direction of $\mathbf{a}_k(t+1)$. After the construction of the exchange estimates, the final estimate $\mathbf{w}_i(t+1)$ is calculated as

$$\mathbf{w}_i(t+1) = \lambda_{i,i} \phi_i(t+1) + \|\phi_i(t+1)\| \sum_{k \in \mathcal{N}_i \setminus i} \lambda_{i,k} \mathbf{a}_k(t+1).$$

Remark 3: Fig. 1 also demonstrates the update procedure for $\mathbf{a}_k(t) = \mathbf{a}_2$. The update is performed if the line, $\mathbf{c}_\perp$, perpendicular to $\mathbf{c}(t+1)$ passes through the shaded update region. Otherwise, the exchanged sign provides no new information and is discarded.

Alternatively, we can also resolve the amplitude uncertainty by using a sign LMS [3] based approach. In this approach, at each node, we run an adaptive algorithm considering $\mathbf{c}(t+1)^T \phi_k(t+1)$ as the desired data and $\mathbf{c}(t+1)$ as the regression vector. We then diffuse the sign of the error $z_k(t+1) = \epsilon_k(t+1) \triangleq \mathbf{c}(t+1)^T \phi_k(t+1) - \mathbf{c}(t+1)^T \mathbf{a}_k(t)$. Using the sign algorithm [3], each node k can construct the exchange estimate as

$$\mathbf{a}_k(t+1) = \mathbf{a}_k(t) + \sigma_k \text{sign}(z_k(t+1)) \mathbf{c}(t+1). \quad (7)$$

Assuming $\mathbf{a}_k(t)$'s are initialized with the same values at each node, (7) can be repeated at all neighboring nodes of k to produce the same $\mathbf{a}_k(t)$. In the next section, we analyze the global stability of the algorithms in the mean.

IV. STABILITY ANALYSIS

We can write the reduced dimension diffusion (2) and the sign algorithm inspired diffusion (7) approaches in a compact form as

$$\phi_i(t+1) = \mathbf{w}_i(t) + \mu_i \mathbf{u}_i(t) e_i(t), \quad (8)$$

$$\mathbf{a}_k(t+1) = \mathbf{a}_k(t) + \sigma_k \mathbf{c}(t+1) h(\epsilon_k(t+1), \mathbf{c}(t+1)) \quad (9)$$

$$\mathbf{w}_i(t+1) = g_i(\phi_i(t+1), \mathbf{a}_k(t+1); k \in \mathcal{N}_i \setminus i), \quad (10)$$

where $\mu_i > 0$ and $\sigma_k > 0$ are the local learning rates, $g_i(\cdot)$ is a combination function such as (3) and

$$e_i(t) = d_i(t) - \mathbf{u}_i(t)^T \mathbf{w}_i(t),$$

$$\epsilon_k(t+1) = \mathbf{c}(t+1)^T (\phi_k(t+1) - \mathbf{a}_k(t))$$

are the estimation and projected reconstruction errors. Here, $h(\epsilon_k(t+1), \mathbf{c}(t+1))$ is a generic function of $\epsilon_k(t+1)$ and $\mathbf{c}(t+1)$, e.g., for the scalar diffusion case $h(\epsilon_k(t+1), \mathbf{c}(t+1)) = (\mathbf{c}(t+1)^T \mathbf{c}(t+1))^{-1} \epsilon_k(t+1)$.

We define deviations from the parameter of interests as

$$\Delta \phi_k(t+1) = \mathbf{w}_o - \phi_k(t+1), \quad (11)$$

$$\Delta \mathbf{a}_k(t+1) = \phi_k(t+1) - \mathbf{a}_k(t+1). \quad (12)$$

Substituting (12) into (10), we get the final estimate as

$$\mathbf{w}_i(t+1) = \sum_{k \in \mathcal{N}_i} \lambda_{i,k} \phi_k(t+1) - \sum_{k \in \mathcal{N}_i \setminus i} \lambda_{i,k} \Delta \mathbf{a}_k(t+1). \quad (13)$$

In the analysis of the mean stability, we make the following assumptions:

- 1) The projection signal $\mathbf{c}(t)$ (or $\mathbf{C}(t)$) and the regression data $\mathbf{u}_k(t)$ are temporally independent.
- 2) The *a priori* construction error $\epsilon_k(\cdot)$ and the projection signal $\mathbf{c}(\cdot)$ (or $\mathbf{C}(\cdot)$) are jointly Gaussian. For sufficiently small step size and long filter length, this assumption is true [3].
- 3) The original parameter estimates $\phi_i(\cdot)$ vary slowly relative to the constructed estimates $\mathbf{a}_i(\cdot)$ through the appropriate step sizes such that

$$\Delta \mathbf{a}_k(t) = \phi_k(t) - \mathbf{a}_k(t) \cong \phi_k(t+1) - \mathbf{a}_k(t) \text{ or}$$

$$\Delta \mathbf{a}_k(t+1) = \phi_k(t+1) - \mathbf{a}_k(t+1) \cong \phi_k(t) - \mathbf{a}_k(t+1).$$

We then define the following global variables

$$\Delta \phi(t) \triangleq \begin{bmatrix} \Delta \phi_1(t) \\ \vdots \\ \Delta \phi_N(t) \end{bmatrix}, \quad \Delta \mathbf{a}(t) \triangleq \begin{bmatrix} \Delta \mathbf{a}_1(t) \\ \vdots \\ \Delta \mathbf{a}_N(t) \end{bmatrix},$$

$$\mathbf{U}(t) \triangleq \begin{bmatrix} \mathbf{u}_1(t) & \dots & \mathbf{0} \\ \vdots & \ddots & \vdots \\ \mathbf{0} & \dots & \mathbf{u}_N(t) \end{bmatrix}, \quad \mathbf{v}(t) \triangleq \begin{bmatrix} v_1(t) \\ \vdots \\ v_N(t) \end{bmatrix},$$

where the vector dimensions are $(mN \times 1)$ and the matrix dimensions are $(mN \times N)$.

Using (8), (9), (11), (12), and (13), we get

$$\Delta \phi(t+1) = (\mathbf{I} - \mathbf{D}\mathbf{U}(t)\mathbf{U}(t)^T) \mathbf{G} \Delta \phi(t) - (\mathbf{I} - \mathbf{D}\mathbf{U}(t)\mathbf{U}(t)^T) \tilde{\mathbf{G}} \Delta \mathbf{a}(t) + \mathbf{D}\mathbf{U}(t)\mathbf{v}(t), \quad (14)$$

where $\mathbf{G} \triangleq \mathbf{\Lambda} \otimes \mathbf{I}_m$ is the transition matrix (and $\otimes$ is the Kronecker product), $\tilde{\mathbf{G}} \triangleq \mathbf{G} - \text{diag}(\mathbf{G})$, $\mathbf{\Lambda} \triangleq [\lambda_{i,k}]$ is the combination matrix and $\mathbf{D} \triangleq \text{diag}([\mu_1, \mu_2, \dots, \mu_N]) \otimes \mathbf{I}_m$.

The global update for the reconstructed parameters yields

$$\Delta \mathbf{a}(t+1) = (\mathbf{I} - \mathbf{S}\mathbf{H}(t)) \Delta \mathbf{a}(t), \quad (15)$$

where $\mathbf{S} \triangleq \text{diag}([\sigma_1, \sigma_2, \dots, \sigma_N]) \otimes \mathbf{I}_m$ and $\mathbf{H}(t)$ is an appropriate transition matrix. As an example, for the scalar diffusion case we have

$$\mathbf{H}(t) = \mathbf{I}_m \otimes \left(\frac{\mathbf{c}(t+1)\mathbf{c}(t+1)^T}{\mathbf{c}(t+1)^T \mathbf{c}(t+1)} \right).$$

For the single-bit diffusion, $h(\epsilon_k(t+1), \mathbf{c}(t+1)) = \text{sign}(\epsilon_k(t+1))$ is a nonlinear function of $\epsilon_k(t+1)$, hence it is not straightforward to write (15). Although $h(\epsilon_k(t+1), \mathbf{c}(t+1))$ is nonlinear, it can be linearized using a Taylor series expansion. However, since we assume joint Gaussianity, by the Price's theorem [3], we can write the expectation of the deviation $\Delta \mathbf{a}_k(t+1)$ as [3]

$$\Delta \bar{\mathbf{a}}_k(t+1) = \Delta \bar{\mathbf{a}}_k(t) - \sigma_k \sqrt{\frac{2}{\pi}} \frac{E[\mathbf{c}(t+1)\mathbf{c}(t+1)^T]}{E[\epsilon_k^2(t+1)]} \Delta \bar{\mathbf{a}}_k(t).$$

Note that the variance $E[\epsilon_k^2(t+1)]$ is given by $E[(\mathbf{c}(t+1)\Delta \bar{\mathbf{a}}(t))^2]$, and we define $\mathbf{F}(t+1) \triangleq \sqrt{\frac{\pi}{2}} \text{diag}([\epsilon_1^2(t+1), \dots, \epsilon_N^2(t+1)])$ that leads

$$\bar{\mathbf{H}}(t) = [\bar{\mathbf{F}}(t+1) \otimes \mathbf{I}_m]^{-1} [\mathbf{I}_m \otimes (E[\mathbf{c}(t+1)\mathbf{c}(t+1)^T])].$$

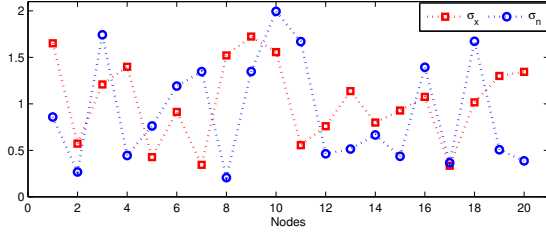


Fig. 2: Statistical profile of the example network.

Taking the expectation of both sides of (14) and (15) and by the first assumption, we get

$$\begin{bmatrix} \Delta \bar{\phi}(t+1) \\ \Delta \bar{a}(t+1) \end{bmatrix} = \begin{bmatrix} BG & B\tilde{G} \\ \mathbf{0} & I - S\bar{H}(t) \end{bmatrix} \begin{bmatrix} \Delta \bar{\phi}(t) \\ \Delta \bar{a}(t) \end{bmatrix}, \quad (16)$$

where $B \triangleq I - \overline{DU(t)U(t)^T}$. (16) covers both the reduced dimension and single bit diffusion strategies. From (16) we also observe that our algorithms are stable in the mean if $|\lambda(I - S\bar{H})| < 1$ (provided that the full diffusion scheme is stable), where $\lambda(\cdot)$'s are the eigenvalues. As example, for the scalar case, assuming $c(\cdot)$ are i.i.d. zero mean with unit variance, then $|\lambda(I - S\bar{H})| < 1$ if and only if $|1 - \sigma_i| < 1$ for all i . Furthermore, the step sizes σ_i for the reconstruction algorithms could be chosen accordingly for comparable convergence performance with the full diffusion case. Following examples illustrate these results.

V. NUMERICAL EXAMPLE AND CONCLUDING REMARKS

In this section, we compare the introduced algorithms with the full diffusion and no-cooperation schemes for the example network with $N = 20$ nodes. Here, we have stationary data $d_i(t) = \mathbf{w}_o^T \mathbf{u}_i(t) + v_i(t)$ for $i = 1, 2, \dots, N$, where $\mathbf{u}_i(t)$ and $v_i(t)$ are i.i.d. zero mean and their variances and $\mathbf{w}_o \in \mathbb{R}^4$ are randomly chosen (See Fig. 2).

The combination matrix $\Lambda = [\lambda_{i,k}]$ is chosen as

$$\lambda_{i,k} = \begin{cases} 1/\max(n_i, n_k) & \text{if } i \neq k \text{ are linked,} \\ 0 & \text{for } i \text{ and } k \text{ not linked,} \\ 1 - \sum_{k \in \mathcal{N}_i \setminus i} \lambda_{i,k} & \text{for } i = k, \end{cases}$$

where n_i and n_k denote the number of neighboring nodes for i and k according to the Metropolis rule.

The step sizes for the adaptation algorithms (8) of diffusion schemes are set such that they converge with the same rate: 0.1 for no-cooperation scheme (i.e. the combination matrix $\Lambda = I_N$), 0.2 for the single-bit and reduced dimension diffusion strategies, and 0.028 for the full diffusion configuration. The step sizes σ_i for the reconstruction algorithms (9) are set as 0.25, 0.72, 0.36, and 0.18 for the single-bit, one-dimension, two-dimension, and three-dimension, respectively. The randomized projection vectors $c(t)$ (and matrices $C(t)$) are generated i.i.d. zero mean Gaussian with standard deviation 1. We point out that we set the learning rates for all algorithms such that the convergence rate of all algorithms are the same for a fair comparison of the final MSDs.

In Fig. 3, we compare the mean-square deviation of various diffusion schemes in terms of their steady state MSDs for the same convergence rate. As expected, in our simulations, the introduced algorithms readily outperform the no-cooperation scheme in terms of the final MSDs. We observe from these

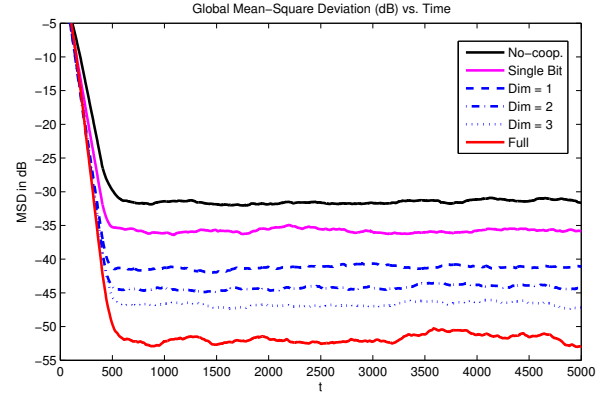


Fig. 3: Global mean-square deviation (MSD) of diffusion and no-cooperation schemes.

simulations that although we significantly reduce the amount of information exchange, the introduced algorithms perform similar to the full information case. To illustrate this further, in Fig. 3, we plot the performance of the reduced-dimension algorithm where we gradually increase the number of dimensions that we kept. We observe that as the number of dimensions increases, the reduced-dimension algorithm gradually achieves the performance of the full information case. Note also that the no-cooperation scheme gives stable error because the adaptation algorithms converge at all nodes. If at least one of the nodes diverges, then the performance of the no-cooperation scheme degrades severely whereas this does not usually influence the diffusion algorithms.

In this letter we introduce novel diffusion based distributed adaptive estimation algorithms that significantly reduce the communication load while providing comparable performance with the full information exchange approaches in our simulations. We achieve this by exchanging either a scalar or a single bit of information generated from random projections of the estimated vectors at each node. Based on these exchanged information, each node recalculates the estimates generated by its neighboring nodes (which are then subsequently merged). We also provide a mean stability analysis of the introduced approaches for stationary data. This analysis can also be extended to mean-square and tracking analysis under certain settings.

REFERENCES

- [1] C. G. Lopes and A. H. Sayed, "Diffusion least-mean squares over adaptive networks: Formulation and performance analysis," *IEEE Transactions on Signal Processing*, vol. 56, no. 7, pp. 3122–3136, 2008.
- [2] F. S. Cattivelli and A. H. Sayed, "Diffusion lms strategies for distributed estimation," *IEEE Transactions on Signal Processing*, vol. 58, no. 3, pp. 1035–1048, 2010.
- [3] A. H. Sayed, *Fundamentals of Adaptive Filtering*. John Wiley and Sons, 2003.
- [4] S. Xie and H. Li, "Distributed LMS estimation over networks with quantised communications," *International Journal of Control*, vol. 86, no. 3, pp. 478–492, 2013.
- [5] A. Ribeiro, G. Giannakis, and S. Roumeliotis, "Soi-kf: Distributed kalman filtering with low-cost communications using the sign of innovations," *Signal Processing, IEEE Transactions on*, vol. 54, no. 12, pp. 4782–4795, 2006.
- [6] H. Sayyadi and M. R. Doostmohammadian, "Finite-time consensus in directed switching network topologies and time-delayed communications," *Scientia Iranica*, vol. 18, no. 1, pp. 75–85, February 2011.
- [7] S. Y. Tu and A. H. Sayed, "On the influence of informed agents on learning and adaptation over networks," *IEEE Transactions on Signal Processing*, vol. 61, no. 6, pp. 1339–1356, March 2013.
- [8] S. S. Kozat, A. T. Erdogan, A. C. Singer, and A. H. Sayed, "Steady state MSE performance analysis of mixture approaches to adaptive filtering," *IEEE Transactions on Signal Processing*, vol. 58, no. 8, pp. 4050–4063, August 2010.